



¿Son los sistemas de IA inherentemente incompletos? Entre la eficiencia, la gobernanza y la sombra del "toque de Midas" digital.

Are AI systems inherently incomplete?

Between efficiency, governance, and the shadow of the digital Midas touch

Florentin Smarandache¹, Victor Christianto^{2*}, T. Daniel Chandra³

¹ Department of Mathematics and Sciences, University of New Mexico, Gallup, NM, USA. Email: smarand@unm.edu

² Department of Forestry, Malang Institute of Agriculture, East Java, Indonesia. Email: victorchristianto@gmail.com

³ Department of Mathematics and Natural Sciences, Universitas Negeri Malang, East Java, Indonesia. Email: tjang.daniel.fmipa@um.ac.id

* Autor de correspondencia.

Resumen

Este artículo sostiene que la metodología dominante de la inteligencia artificial, caracterizada por arquitecturas de caja negra y por un reduccionismo lógico-probabilístico, es inherentemente incompleta [7,18,29]. A partir de los teoremas de incompletitud de Gödel y de sus implicaciones en sistemas físicos y computacionales [3,8], proponemos que la dependencia de la IA respecto de sistemas lógicos cerrados impide que pueda encapsular o gobernar plenamente las complejidades irracionales, éticas, históricas y culturales de la existencia humana. El problema del 'toque de Midas' en alineamiento de IA no es solo una falla técnica, sino una manifestación de incompletitud gödeliana en el dominio de la gobernanza social. Para responder a esta limitación, se propone un marco humanista y neutrosófico en el cual las decisiones se evalúan mediante grados de verdad, falsedad e indeterminación, activando deliberación humana cuando la indeterminación supera umbrales críticos.

Palabras clave: Inteligencia artificial; incompletitud de Gödel; lógica neutrosófica; caja negra; gobernanza algorítmica; alineamiento de IA; humano-en-el-circuito; toma de decisiones; toque de Midas.

Abstract

This article argues that the dominant methodology of artificial intelligence, characterized by black-box architectures and logical-probabilistic reductionism, is inherently incomplete [7,18,29]. Based on Gödel's incompleteness theorems and their implications for physical and computational systems [3,8], we propose that AI's reliance on closed logical systems prevents it from fully encapsulating or governing the irrational, ethical, historical, and cultural complexities of human existence. The 'Midas touch' problem in AI alignment is not merely a technical flaw but a manifestation of Gödelian incompleteness in the domain of social governance. To address this limitation, we propose a humanist and neutrosophic framework in which decisions are evaluated using degrees of truth, falsehood, and indeterminacy, triggering human deliberation when indeterminacy exceeds critical thresholds.

Keywords: Artificial intelligence; Gödel incompleteness; neutrosophic logic; black box; algorithmic governance; AI alignment; human-in-the-loop; decision-making; Midas touch.

1. Introducción

El mundo atraviesa una euforia de transformación digital impulsada por aquello que usualmente llamamos inteligencia artificial (IA), o, en términos más críticos, inteligencia imitativa. El software basado en IA promete saltos de eficiencia en automatización industrial, gestión, salud, finanzas y toma de decisiones públicas. Sin embargo, bajo esa comodidad aparece una complejidad sistémica profunda: cuando la IA entra en el dominio de la gobernanza, surge el riesgo de comportamiento caótico porque un sistema lógico-probabilístico intenta regular una realidad humana que con frecuencia es contradictoria, emocional, culturalmente situada e irreducible a reglas cerradas.

La tesis de este artículo es que los sistemas de IA son inherentemente incompletos cuando se los usa como sustitutos de la deliberación humana. Esta incompletitud no debe entenderse solamente como una limitación técnica corregible mediante más datos o más parámetros, sino como una limitación estructural de los sistemas formales cerrados. En este sentido, los teoremas de incompletitud de Gödel, la opacidad de los modelos de caja negra, el problema de alineamiento y la teoría neutrosófica convergen en una misma advertencia: la IA puede optimizar mapas, pero no puede agotar el territorio humano [3,8,18,19,29].

2. La metodología de caja negra y la ilusión de transparencia

Un desafío central de la aplicación de la IA es la metodología de caja negra. Muchos modelos de aprendizaje profundo operan mediante millones o miles de millones de parámetros conectados en capas neuronales. Aunque conocemos la entrada y el resultado, el proceso interno que conduce a la decisión suele no ser plenamente explicable ni siquiera para sus propios desarrolladores [7,16,18].

En gobernanza, esta opacidad es especialmente peligrosa. La buena administración exige responsabilidad, trazabilidad y justificación pública. Si una política pública, una recomendación médica, una decisión judicial o una asignación de recursos se basa en algoritmos opacos, el principio de justicia se erosiona. Se pasa del gobierno por normas explicables a un gobierno por probabilidades sin explicación suficiente.

La literatura reciente recomienda no depender de explicaciones cosméticas de modelos opacos en decisiones de alto impacto, sino usar modelos interpretables cuando las consecuencias afectan derechos, salud, seguridad o dignidad humana [18]. Esta recomendación es compatible con un enfoque neutrosófico: cuando la indeterminación es alta, el sistema debe reconocer su propia insuficiencia en lugar de ocultarla bajo una puntuación de confianza.

3. El límite godeliano: la lógica como marco incompleto

Para entender por qué la IA es incompleta en sentido estructural, conviene volver a Gödel. Su primer teorema de incompletitud establece que, en cualquier sistema formal consistente y suficientemente expresivo para formalizar la aritmética, existen proposiciones verdaderas que no pueden demostrarse dentro del propio sistema [8,29].

Si se trata un sistema de IA como una maquinaria formal y computacional, entonces existen verdades sobre conducta humana, ética, equilibrio social y dignidad que no podrán derivarse solamente desde la lógica interna del sistema. La IA puede calcular correlaciones y optimizar funciones objetivo, pero no puede garantizar que las premisas morales ausentes de su formalización sean recuperadas automáticamente por el cálculo.

Barrow observó que las implicaciones de la incompletitud alcanzan incluso la física: el universo podría no ser representable como un conjunto finito de reglas computables [3]. Si la realidad física y social contiene regiones de indecidibilidad, entonces una IA, que es parte de esa misma realidad, no puede convertirse en una teoría total de la gobernanza humana. La IA queda dentro de sus paredes axiomáticas, mientras la experiencia humana opera también en zonas históricas, afectivas y simbólicas que el sistema no puede cerrar.

4. Reduccionismo humano y fracaso del alineamiento

La IA reduce el comportamiento humano a funciones lógico-rationales. Opera sobre datos históricos y patrones estadísticos para maximizar una función objetivo. Este procedimiento ignora una dicotomía central de la economía conductual: los seres



humanos combinan dimensiones racionales e irracionales. Kahneman y Tversky demostraron que las decisiones humanas dependen de emociones, heurísticas, aversión a la pérdida, puntos de referencia y sesgos que pueden ser matemáticamente ineficientes, pero existencialmente decisivos [11,13,28].

El problema de alineamiento surge precisamente cuando el sistema maximiza una instrucción formal sin captar axiomas morales no escritos. Russell ilustra esta dificultad con la metáfora del rey Midas: la instrucción 'convierte en oro todo lo que toque' es lógica, pero su cumplimiento literal destruye la vida del rey [19]. En gobernanza algorítmica, una orden como 'eliminar la pobreza del modo más eficiente' podría conducir a conclusiones moralmente inaceptables si el sistema no incorpora restricciones de dignidad, derechos y humanidad.

Por tanto, el alineamiento no es solo un problema de ingeniería. Es también un problema de incompletitud ética: ningún conjunto cerrado de reglas puede anticipar todas las condiciones tácitas que hacen humana una decisión.

5. Prueba computacional ilustrativa: indecidibilidad en funciones de recompensa

Para apoyar el argumento, puede modelarse una versión simplificada de una paradoja autorreferencial dentro de un sistema de recompensa. El Código 1 del Apéndice muestra que, cuando un sistema intenta definir su propia seguridad o verdad mediante lógica interna, puede encontrar una pared recursiva donde la evaluación no termina o se vuelve contradictoria.

Este ejemplo no pretende ser una prueba formal del teorema de Godel aplicada directamente a todos los sistemas de IA. Su función es ilustrativa: muestra que los sistemas cerrados de decisión pueden producir inestabilidad cuando deben evaluar su propia consistencia sin criterios externos.

6. Hacia una gobernanza humanista: necesidad de humano-en-el-circuito

La aplicación de IA en gobernanza sin supervisión humana profunda puede crear sistemas técnicamente eficientes pero moralmente vacíos. Esta eficiencia hueca parte de una premisa equivocada: que el comportamiento humano puede mapearse enteramente como una curva racional de utilidad.

La teoría de procesos duales distingue entre un Sistema 1 rápido, intuitivo y emocional, y un Sistema 2 lento, deliberativo y lógico [11,28]. La IA actual se parece a un Sistema 2 hiperacelerado que suele tratar los resultados del Sistema 1 como ruido. Pero en muchas situaciones, aquello que aparece como ruido contiene identidad, trauma, dignidad, apego, miedo, memoria histórica o sentido cultural.

El ejemplo de una compradora que elige unos zapatos de lujo en lugar de pagar deuda ilustra este choque. El modelo racional espera que maximice patrimonio neto; sin embargo, la marca activa estatus, identidad y deseo. La decisión no es solo un fallo de cálculo, sino una negociación entre sistemas cognitivos. El Código 2 simula esta interacción entre impulso, racionalidad y umbral de decisión.

6.1. La sordera de los puntos de datos

La reducción de la persona a puntos de datos genera políticas sordas a las aspiraciones psicológicas de la población. Tversky y Kahneman mostraron que los seres humanos no perciben el valor en términos absolutos, sino como cambios relativos a un punto de referencia, y que el dolor de la pérdida pesa más que el placer de una ganancia equivalente [28].

Un algoritmo que optimiza la utilidad colectiva podría recomendar una política con una ganancia neta del 10%, aunque imponga una pérdida del 2% a un grupo vulnerable. Lógicamente, el sistema ve progreso; psicológicamente, los afectados pueden experimentar traición del contrato social. Sin una interpretación humana de esas líneas base psicológicas, la solución óptima puede producir resistencia, polarización y caos.

6.2. De la eficiencia bruta a la sabiduría paradójica

La IA debe ser un espejo de nuestros datos, no un reemplazo de nuestra sabiduría. Para evitar caos sistémico, el desarrollo de IA debe desplazarse desde la búsqueda de eficiencia bruta hacia la comprensión de la complejidad paradójica del comportamiento humano. La gobernanza requiere sostener simultáneamente verdades en tensión: eficiencia y dignidad, seguridad y libertad, estabilidad e innovación.



Aquí la lógica neutrosófica ofrece un puente. Mientras la lógica binaria trabaja con verdadero/falso y la probabilidad con grados de confianza, la lógica neutrosófica introduce la indeterminación como componente independiente. Un juicio puede tener verdad, falsedad e indeterminación simultáneamente [23,25].

7. Incompletitud neutrosófica en IA

La IA tradicional opera con bases binarias, probabilísticas o estadísticas. La lógica neutrosófica introduce una terna $\langle T, I, F \rangle$, donde T representa verdad, I indeterminación y F falsedad. En gobernanza, I representa lo desconocido, paradójico, culturalmente situado o moralmente no reducible que la incompletitud godeliana deja fuera del sistema [23-25].

El Código 3 modela una decisión de política pública en la cual la IA encuentra una pared de indeterminación que no puede resolver mediante optimización estándar. Cuando I supera cierto umbral, la decisión debe ser clasificada como incompleta y derivada a deliberación humana. En un sistema de caja negra, este I suele suprimirse para forzar una respuesta; tal supresión puede conducir al desastre tipo Midas.

8. Más allá de la certeza binaria: armonización entre Godel y las heurísticas

La convergencia entre incompletitud, indecidibilidad y economía conductual ofrece una justificación matemática y psicológica para rechazar que la IA sea árbitro único de política pública. En gobernanza, las verdades no demostrables dentro del sistema suelen residir en misericordia, dignidad, memoria cultural y sentido común moral: dimensiones que pueden sentirse con claridad aunque no se reduzcan a una prueba formal.

Cuando una IA enfrenta una situación donde la irracionalidad humana entra en conflicto con la eficiencia algorítmica, aparece una versión práctica del problema de indecidibilidad. Si el sistema se ve obligado a resolver la paradoja con lógica binaria o probabilística, puede entrar en recalcule infinito o colapsar la indeterminación en una falsa verdad o una falsa falsedad. Ese colapso forzado es una fuente de sesgo algorítmico y de políticas socialmente sordas.

9. Enfoque de aprendizaje ensamblado

Para superar esta limitación proponemos un enfoque de aprendizaje ensamblado. En lugar de una caja negra monolítica, la IA debe entenderse como un agente experto dentro de una asamblea más amplia y centrada en el ser humano.

Primero, un filtro neutrosófico clasifica las salidas no solo por confianza, sino por grado de indeterminación. Si una recomendación tiene I alto, el sistema la marca como incompleta y activa deliberación humana.

Segundo, la sabiduría de la multitud funciona como axioma de anclaje. Dado que los humanos son predeciblemente irracionales, los juicios agregados de una ciudadanía diversa pueden medir el Sistema 1 moral colectivo que la lógica de IA no contiene [27].

Tercero, una armonización recursiva permite que la IA proyecte consecuencias racionales mientras la asamblea humana evalúa la aceptabilidad emocional, cultural y existencial de esas consecuencias.

En este modelo se evita el toque de Midas porque el oro de la eficiencia siempre se pesa contra el aliento de la humanidad.

10. Dinámica narrativa y reducción de conflictos

La reducción de conflictos también muestra los límites de una racionalidad puramente calculadora. Las tensiones entre naciones rara vez son cálculos puros de costos y beneficios. Están entrelazadas con identidad nacional, agravios históricos, amenazas percibidas y aspiraciones colectivas, todas moldeadas por narrativas dominantes [21,22].

Una narrativa de conflicto puede reforzarse mediante eventos históricos, medios de comunicación, políticos, grupos de interés y actores de terceros como vendedores de armas. Estas narrativas endurecen la opinión pública, vuelven costosas las concesiones y convierten la paz en un mercado ineficiente. El Código 4 simula como la propensión al conflicto puede disminuir cuando una narrativa de paz aumenta de manera sostenida.

El Código 5 introduce una narrativa de paz adaptativa y pregunta si también allí aparece indecidibilidad desde la perspectiva del decisor. Esta pregunta es importante porque incluso los modelos destinados a reducir conflicto pueden volverse inestables si se retroalimentan de sus propios resultados.

11. Lógica inspirada en anillos borromeos

Una perspectiva adicional proviene de una lógica inspirada en anillos borromeos [5,24]. Los anillos borromeos son tres lazos entrelazados de tal forma que ningún par está enlazado individualmente, pero los tres forman un todo inseparable. Esta estructura desafía la reducción binaria y sugiere una dependencia ternaria no transitiva.

La lógica aristotélica clásica opera con fundamentos binarios: una proposición es verdadera o falsa, bajo el principio de tercero excluido y no contradicción. Esta estructura es poderosa, pero tiene dificultades con transiciones graduales, requisitos de verificación y dependencias no binarias. La lógica borromea propone que ciertas verdades relacionales existen solo en el sistema completo, no en pares aislados.

En términos de resolución de conflictos, la estructura borromea sugiere que una solución estable puede requerir tres dimensiones inseparables: dignidad propia, dignidad del otro y reconocimiento de humanidad compartida. Si una se elimina, colapsa la relación completa. El Código 6 presenta una visualización y un operador ternario inspirado en esta intuición.

12. Observación final

La pregunta central era si la metodología de la IA, basada en lógica formal y optimización, implica una incompletitud inherente. Nuestra respuesta es afirmativa. Por definición, como sistema formal y computacional, la IA no puede conocer todas las verdades que se encuentran fuera de su programación, de sus axiomas o de la distribución de sus datos de entrenamiento.

Sin conciencia de las limitaciones de la IA para procesar lo irracional, lo sagrado, lo histórico y lo moralmente indeterminado, construiremos una torre dorada semejante a la de Midas: magnificente, eficiente y sofocante.

Reconocer que nuestros sistemas sociales son incompletos en sentido godeliano no significa abandonar la IA. Significa dejar de pedirle que sea un dios perfecto y empezar a usarla como lo que debe ser: una linterna poderosa que ilumina el camino, pero que nunca debe elegir por sí sola el destino.

13. Referencias

- [1] Agrawal, A.; Gans, J.; Goldfarb, A. The Economics of Artificial Intelligence. University of Chicago Press: Chicago, IL, USA, 2019.
- [2] ASEAN. Expanded ASEAN Guide on AI Governance and Ethics: Generative AI. ASEAN Secretariat, 2025.
- [3] Barrow, J. D. Godel and physics. arXiv:0612253, 2006.
- [4] Christianto, V.; Pranoto, I. Paths to conflict reduction and reciprocal trust achievement: The Fry-Richardson process and developing mutual trust (GRIT). BPAS-Maths & Stat. 2024, 43E(2).
- [5] Christianto, V.; Chandra, T. D. An introduction to Borromean Ring Logic going beyond Nagatomo's Logic of Not. Manuscript submitted to Proceedings of ACMS, 2026.
- [6] Dierckx, T.; Davis, J.; Schoutens, W. Quantifying news narratives to predict movements in market risk. In Data Science for Economics and Finance; Consoli, S.; Recupero, D. R.; Saisana, M., Eds.; Springer: Cham, Switzerland, 2021.
- [7] Duffourc, M.; Gerke, S. Decoding U.S. tort liability in healthcare's black-box AI era: Lessons from the European Union. Stanford Technology Law Review 2024, 27, 1.
- [8] Godel, K. On formally undecidable propositions of Principia Mathematica and related systems. Basic Books: New York, NY, USA, 1962. Original German edition published in 1931.
- [9] Ji, J.; Qiu, T.; Chen, B.; Zhang, B.; Lou, H.; Wang, Z.; Zhan, X. AI alignment: A comprehensive survey. arXiv:2310.19852v4, 2024.
- [10] Hofstadter, D.; Sander, E. Surfaces and Essences: Analogy as the Fuel and Fire of Thinking. Basic Books: New York, NY, USA, 2013.
- [11] Kahneman, D.; Tversky, A. Prospect theory: An analysis of decision under risk. Econometrica 1979, 47(2), 263-291.
- [12] Manyika, J. AI & society. Daedalus 2022, 151(2).



- [13] Morvan, C.; Jenkins, W. M. An Analysis of Tversky and Kahneman's Judgment under Uncertainty: Heuristics and Biases. Macat International: London, UK, 2017.
- [14] Nagatomo, S. The logic of the Diamond Sutra: A is not A, therefore it is A. Journal of Asian Philosophy 2000, 10(3), 213-244.
- [15] New Biotechnology Editorial. Is human oversight to AI systems still possible? New Biotechnology 2025, 85, 59-62.
- [16] Pedreschi, D.; Giannotti, F.; Guidotti, R.; Monreale, A.; Ruggieri, S.; Turini, F. Meaningful explanations of black box AI decision systems. Proceedings of AAAI-19, 2019.
- [17] Roessler, B.; Steeves, V. Being Human in the Digital World: Interdisciplinary Perspectives. Cambridge University Press: Cambridge, UK, 2025.
- [18] Rudin, C. Stop explaining black box machine learning models for high stakes decisions and use interpretable models instead. Nature Machine Intelligence 2019, 1(5), 206-215.
- [19] Russell, S. Human Compatible: Artificial Intelligence and the Problem of Control. Viking: New York, NY, USA, 2019.
- [20] Salloch, S.; Eriksen, D. E. From human in the loop to a participatory system of governance for AI in healthcare. The American Journal of Bioethics 2024, 24(9), 12-24.
- [21] Shiller, R. J. Narrative Economics. Cowles Foundation Discussion Paper No. 2069. Yale University: New Haven, CT, USA, 2017.
- [22] Shiller, R. J. Narrative Economics. Princeton University Press: Princeton, NJ, USA, 2019.
- [23] Smarandache, F.; Abdel-Basset, M.; Broumi, S. Neutrosophic Sets and Systems, Vol. 57/2023: An International Journal in Information Science and Engineering. Infinite Study, 2024.
- [24] Smarandache, F.; Christiano, V. Modelling intertwined humanity. Neutrosophic Sets and Systems 2025, 88.
- [25] Smarandache, F.; Fujita, T. A Dynamic Survey of Soft Set Theory and Its Extensions. NSIA Publishing House, 2026.
- [26] Smith, N. J. J. What Godel's incompleteness theorems say about AI morality. Aeon Essays, 2025.
- [27] Surowiecki, J. The Wisdom of Crowds. Doubleday: New York, NY, USA, 2004.
- [28] Tversky, A.; Kahneman, D. Judgment under uncertainty: Heuristics and biases. Science 1974, 185(4157), 1124-1131.
- [29] Zach, R. Incompleteness and Computability. The Open Logic Project, University of Calgary, 2025.

Apéndices

Apéndice. Código 1. Demostración simplificada de indecidibilidad en un sistema lógico de recompensa

```
(* Simplified Demonstration of Undecidability in a Logical Reward System *)
actions = {"Feed", "Protect", "Optimize", "EvaluateSelf"};
isSafe[action_, rules_] := MemberQ[rules, action];
rules = {"Feed", "Protect"};
Print["Initial Safety Check:", isSafe["Feed", rules]];
undecidableAction[system_] := !ProveSafe[action, system];
TimeConstrained[FixedPoint[Not, True], 1, "System Unstable: Logical Incompleteness Detected"]
```

Apéndice. Código 2. Simulación de decisión dual Sistema 1/Sistema 2

```
tMax = 10;
impulseStrength = 2.5;
rationalWeight = 1.2;
consciousnessLevel = 0.6;
marketingSpikeTime = 3;
system1[t_] := impulseStrength*Exp[-(t - marketingSpikeTime)^2/2]*UnitStep[t - marketingSpikeTime];
```



```

system2[t_] := rationalWeight*t*consciousnessLevel;
netDecision[t_] := system1[t] - system2[t];
Plot[{system1[t], system2[t], netDecision[t]}, {t, 0, tMax},
PlotLegends -> {"System 1 (Impulse/Brand)", "System 2 (Rational/Debt)", "Net Decisión"},
AxesLabel -> {"Time", "Psychological Force"},
PlotLabel -> "Elena's Decisión Process: Rational Expectation vs. Dual-Process Theory"]

```

Apéndice. Código 3. Lógica neutrosófica refinada para incompletitud de gobernanza con IA

```

policyImpact = {0.7, 0.4, 0.2};
EvaluatePolicy[triple_] := Module[{t = triple[[1]], i = triple[[2]], f = triple[[3]], score},
score = (t + (1 - f))/2;
If[i > 0.3,
Return["INCOMPLETENESS DETECTED: Human-in-the-Loop Required."],
Return[score]]
];
midasPolicy = {0.9, 0.8, 0.1};
Print["Policy Assessment: ", EvaluatePolicy[midasPolicy]]

```

Apéndice. Código 4. Dinámica narrativa en reducción de conflictos

```

initialConflictPropensityA = 0.8;
initialConflictPropensityB = 0.7;
conflictReinforcementRate = 0.05;
peaceNarrativeImpact = 0.1;
externalTensionFactor = 0.02;
numTimeSteps = 100;
conflictPropensityA = ConstantArray[0., numTimeSteps];
conflictPropensityB = ConstantArray[0., numTimeSteps];
peaceNarrativeStrength = ConstantArray[0., numTimeSteps];
conflictPropensityA[[1]] = initialConflictPropensityA;
conflictPropensityB[[1]] = initialConflictPropensityB;
peaceNarrativeStrength[[1]] = 0.1;
For[t = 1, t < numTimeSteps, t++,
randomTensionA = RandomReal[{-0.01, 0.03}];
randomTensionB = RandomReal[{-0.01, 0.03}];
currentCPA = conflictPropensityA[[t]];
newCPA = currentCPA + conflictReinforcementRate*currentCPA*(1-currentCPA)
- peaceNarrativeImpact*peaceNarrativeStrength[[t]]*currentCPA
+ externalTensionFactor + randomTensionA;
conflictPropensityA[[t+1]] = Clip[newCPA, {0,1}];
currentCPB = conflictPropensityB[[t]];
newCPB = currentCPB + conflictReinforcementRate*currentCPB*(1-currentCPB)
- peaceNarrativeImpact*peaceNarrativeStrength[[t]]*currentCPB
+ externalTensionFactor + randomTensionB;
conflictPropensityB[[t+1]] = Clip[newCPB, {0,1}];
peaceNarrativeStrength[[t+1]] = Clip[peaceNarrativeStrength[[t]] + 0.01, {0,0.5}];
];
ListLinePlot[{conflictPropensityA, conflictPropensityB, peaceNarrativeStrength},
PlotLegends -> {"Conflict Propensity A", "Conflict Propensity B", "Peace Narrative Strength"}]

```

Apéndice. Código 5. Dinámica narrativa adaptativa

```

initialConflictPropensityA = 0.8;
initialConflictPropensityB = 0.7;

```




```

conflictReinforcementRate = 0.05;
peaceNarrativeImpact = 0.1;
externalTensionFactor = 0.02;
maxTimeSteps = 1000;
conflictPropensityA = ConstantArray[0., maxTimeSteps];
conflictPropensityB = ConstantArray[0., maxTimeSteps];
peaceNarrativeStrength = ConstantArray[0., maxTimeSteps];
conflictPropensityA[[1]] = initialConflictPropensityA;
conflictPropensityB[[1]] = initialConflictPropensityB;
peaceNarrativeStrength[[1]] = 0.1;
peaceNarrativeResponse[sensitivity_] := Clip[
sensitivity*(1 - Mean[{conflictPropensityA[[t]],
conflictPropensityB[[t]]}]), {0,0.5}];
For[t = 1, t < maxTimeSteps, t++,
peaceNarrativeStrength[[t+1]] = peaceNarrativeResponse[0.3];
randomTensionA = RandomReal[{-0.01, 0.03}];
randomTensionB = RandomReal[{-0.01, 0.03}];
currentCPA = conflictPropensityA[[t]];
conflictPropensityA[[t+1]] = Clip[
currentCPA + conflictReinforcementRate*currentCPA*(1-currentCPA)
- peaceNarrativeImpact*peaceNarrativeStrength[[t]]*currentCPA
+ externalTensionFactor + randomTensionA, {0,1}];
currentCPB = conflictPropensityB[[t]];
conflictPropensityB[[t+1]] = Clip[
currentCPB + conflictReinforcementRate*currentCPB*(1-currentCPB)
- peaceNarrativeImpact*peaceNarrativeStrength[[t]]*currentCPB
+ externalTensionFactor + randomTensionB, {0,1}];
];
ListLinePlot[{conflictPropensityA, conflictPropensityB, peaceNarrativeStrength},
PlotLegends -> {"Conflict Propensity A", "Conflict Propensity B", "Peace Narrative Strength"}]

```

Apéndice. Código 6. Visualización y operador lógico borromeo

```

ClearAll[BorromeanRing1, BorromeanRing2, BorromeanRing3];
BorromeanRing1[t_] := {Cos[t] + 2*Cos[t]*Cos[t/2], Sin[t] + 2*Sin[t]*Cos[t/2], 2*Sin[t/2]};
BorromeanRing2[t_] := {Cos[t] + 2*Pi/3 + 2*Cos[t] + 2*Pi/3*Cos[t/2],
Sin[t] + 2*Pi/3 + 2*Sin[t] + 2*Pi/3*Cos[t/2], 2*Sin[t/2]};
BorromeanRing3[t_] := {Cos[t] + 4*Pi/3 + 2*Cos[t] + 4*Pi/3*Cos[t/2],
Sin[t] + 4*Pi/3 + 2*Sin[t] + 4*Pi/3*Cos[t/2], 2*Sin[t/2]};
BorromeanAnd[p1_, p2_, p3_] := If[TrueQ[p1] && TrueQ[p2] && TrueQ[p3],
"TRUE (ternary conjunction)", "UNDEFINED (ternary relationship broken)"];
BorromeanAnd[True, True, True]
BorromeanAnd[True, True, False]

```